

Lärdomar från att försöka argumentera med en AI-skeptiker

Episoden började med en politisk meme jag postade: Donald Trump och Benjamin Netanyahu i orange fängelseoveraller, sittande på en våningssäng under ett varmt, nostalgiskt juloverlay med texten "All I Want for Christmas". Den visuella ironin var omedelbar och skarp. Att skapa den krävde medvetna omvägar. Nutida bildgenereringsmodeller har både policybegränsningar och tekniska koherensbegränsningar:

- Grok tillåter karikatyrier av framträdande personer men misslyckas konsekvent med att producera pålitlig överlagrad text.
- ChatGPT är utmärkt på att generera dekorativ festlig text som "All I Want for Christmas" men dess skyddsregler vägrar prompts som avbildar levande politiska ledare i fängelsemiljöer.

Ingen enskild modell kunde producera den kompletta bilden. De motsägelsefulla elementen — laddad politisk satir kombinerad med sentimental julstämning — utlöser vägransmekanismer eller koherensfel. Stora språkmodeller (LLM:er) är helt enkelt oförmögna att syntetisera sådana konceptuellt motsatta komponenter i en sammanhängande utdata. Jag genererade de två elementen separat och slog sedan ihop och redigerade dem manuellt i GIMP. Den slutliga kompositen var otvetydigt mänskligt skapad: min idé, mitt urval av komponenter, min sammanfogning och justeringar. Utan dessa verktyg hade satiren förblivit instängd i mitt huvud eller kommit ut som grova streckgubbar — berövad all visuell effekt.

Någon anmälde bilden som "AI-genererad". Nästa dag införde servern en ny regel som förbjuder generativ AI-innehåll. Denna regel — och memen som utlöste den — inspirerade mig direkt att skriva och publicera essän "Högdimensionella sinnen och serialiseringsbördan: Varför LLM:er är viktiga för neurodivergent kommunikation". Jag hoppades att den skulle uppmuntra till reflektion över hur dessa verktyg fungerar som kognitiva och kreativa stöd. Men det blev istället ett ganska obekvämt utbyte med administratören.

Skeptikerns ståndpunkt och utbytet

Administratören argumenterade för att LLM:er inte utvecklas för människors bästa utan främjar resursslöseri och militarisering. Han nämnde energiförbrukning, militära kopplingar, modellkollaps, hallucinationer och risken för ett "dött internet". Han avslöjade att han bara hade skummat essän och erkände att han äger en kraftfull speldator som kan köra avancerade lokala LLM:er för privat nöje, med tillgång till ännu större modeller via en vän.

Flera motsägelser framträdde:

- Mitt arbete sker på en lågenergidator, en reparerbar Raspberry Pi 5 (5–15 W) med delade molninstanser. Hans lokala setup förbrukar betydligt mer dedikerad energi och hårdvara.
- Hårdvaran han använder för att "leka" med kraftfulla LLM:er lokalt kommer från företag (Intel, AMD, NVIDIA) med miljarder i direkta kontrakt med USA:s försvarsdepartement (DoD).

Mest slående var att personen som införde förbudet för att skydda autenticitet avfärdade någon som aktivt stresstestar LLM:er för faktiska och geopolitiska bias (se mina offentliga granskningar av Grok och ChatGPT).

Hawking-analogin och administratörens egna ord

Administratören identifierade sig som neurodivergent och erkände potentialen i AI som hjälpmedelsteknik. Han berömde realtidsundertextningsglasögon för synskadade som "verkligt coola", men insisterade på att "att låta en maskin skriva essäer och rita bilder är något annat". Han tillade: "Neurodivergenta personer kan göra dessa saker, många har övervunnit hinder för att utveckla dessa färdigheter." Han beskrev också sin egen erfarenhet med LLM:er: "Ju mer jag redan vet om ett ämne, desto mindre behöver jag AI. Ju mindre jag vet om ett ämne, desto mindre rustad är jag att märka hallucinationer och korrigera dem." Dessa uttalanden avslöjar en djup asymmetri i hur stöd bedöms.

Tänk dig att tillämpa samma logik på Stephen Hawking:

"Vi inser att en röstsyntetiserare skulle kunna hjälpa dig att kommunicera snabbare, men vi skulle föredra att du anstränger dig mer med din naturliga röst. Många personer med motorneuronsjukdom har övervunnit hinder för att tala tydligt — du borde utveckla de färdigheterna också. Maskinen gör något annat än riktigt tal."

Eller, ur hans egen syn på faktisk noggrannhet:

"Ju mer Hawking redan vet om kosmologi, desto mindre behöver han syntetiseraren. Ju mindre han vet, desto mindre rustad är han att märka fel i maskinrösten och korrigera dem."

Ingen skulle acceptera detta. Vi förstod att Hawkings syntetiserare inte var en krycka eller utspädning — den var den essentiella bron som tillät hans extraordinära sinne att dela sin fulla djup utan oöverstigliga fysiska hinder.

Administratörens bekvämlighet med linjär, mänskligt strukturerad prosa speglar en kognitiv stil som stämmer bättre överens med neurotypiska förväntningar. Min profil är den omvända: faktisk och logisk djup kommer naturligt (som att utveckla en flerspråkig publiceringsplattform helt själv), men att producera strukturerad, tillgänglig prosa för mänskliga mottagare har alltid varit hindret — exakt vad essän beskriver. Att acceptera undertextningsglasögon eller alt-text som legitima stöd medan man avvisar LLM-strukturering för kognitiv avvikelse är att dra en godtycklig gräns. Mastodon och den bredare Fediversen

stoltserar ofta med sin inkludering. Ändå inför detta nya portar: vissa stöd välkomnas; andra måste övervinnas genom individuell ansträngning.

Historiska ekon: Motstånd mot transformerande verktyg

Det totala avvisandet av offentlig användning av generativ AI ekar ett återkommande mönster genom teknologihistorien. I tidiga 1800-talets England krossade skickliga vävare, kända som ludditer, mekaniserade vävstolar som hotade deras hantverk och försörjning. Gaslykttändare i städer motsatte sig Edisons glödlampa av rädsla för att bli överflödiga. KuskAR, stallarbetare och hästuppfödare motstod bilen som ett existentiellt hot mot deras livsstil. Professionella skrivare och tecknare betraktade fotokopiatorn med oro, eftersom de trodde att den skulle devalvera noggrant handarbete. Sättare och tryckare kämpade mot datoriserad sättning.

I varje fall kom motståndet från genuin rädsla: ny teknik gjorde de färdigheter de var stolta över föråldrade och utmanade deras ekonomiska roller och sociala identitet. Förändringarna kändes som en devalvering av mänskligt arbete.

Ändå bedömer historien dessa innovationer efter deras bredare inverkan: mekanisering minskade slit och möjliggjorde massproduktion; elektrisk belysning förlängde produktiva timmar och förbättrade säkerheten; bilar gav personlig mobilitet; fotokopiatorer demokratiserade tillgång till information; digital sättning gjorde publicering snabbare och mer tillgänglig. Få idag skulle återgå till gaslyktor eller hästdragna transporter bara för att bevara traditionella jobb. Verktygen utökade mänsklig kapacitet och delaktighet långt mer än de minskade den.

Generativ AI — använd som en protes för kognition eller kreativitet — följer samma bana: den utrotar inte mänsklig avsikt utan utökar uttryck för dem vars idéer har begränsats av utförandehinder. Att avvisa den rakt av riskerar att upprepa det ludditiska impulsen — att försvara bekanta processer på bekostnad av bredare delaktighet.

Slutsats: Vem beslutar vilka stöd som är acceptabla?

Händelserna som berättas i denna essä — en anmäld bild, ett hastigt infördt förbud, en utdragen debatt — avslöjar mer än en lokal oenighet om teknik. De blottlägger en långt djupare och mer grundläggande fråga: **Vem får besluta vilka stöd som är acceptabla, och vilka som inte är?** Borde det vara de personer som lever i den hud och hjärna som behöver stödet — de som vet, från daglig erfarenhet, vad som överbryggar klyftan mellan deras förmågor och full delaktighet? Eller borde det vara utomstående, hur välmenande de än är, som inte delar den levda verkligheten och därför inte kan känna tyngden av hindret?

Historien besvarar denna fråga upprepade gånger, och nästan alltid i samma riktning. Rullstolar kritiserades en gång för att uppmuntra beroende; dövundervisningssystem insisterade länge på att barn skulle lära sig läppavläsning och tal istället för teckenspråk. I varje fall segrade till slut de personer som var närmast funktionsnedsättningen — inte för

att de förnekade oro för kostnad, tillgång eller potentiellt missbruk, utan för att de var de primära auktoriteterna på vad som faktiskt återställde deras handlingsfrihet och värdighet.

Med stora språkmodeller och andra generativa verktyg lever vi igenom samma cykel igen. Många som portvaktar deras användning upplever inte de specifika kognitiva eller expressiva hinder som gör linjär strukturering, narrativ flöde eller snabb serialisering till en utmattande översättning till ett främmande språk. Utifrån kan "ansträng dig bara hårdare" eller "utveckla färdigheten" låta rimligt. Inifrån är verktyget inte en genväg runt ansträngning; det är rampen, hörapparaten, protesen som slutligen låter befintlig ansträngning nå världen.

Den djupaste ironin uppstår när beslutsfattarna själva identifierar sig som neurodivergenta, men deras särskilda neurologi stämmer närmare neurotypiska förväntningar i det bedömda området. "Jag övervann det på det här sättet, så andra borde också" är förståeligt, men det fungerar ändå som portvaktning — det reproducerar de normer vi kritiserar när de kommer från neurotypiska auktoriteter. En konsekvent etisk princip behövs:

- Personen närmast funktionsnedsättningen är den primära auktoriteten på vad som möjliggör deras meningsfulla delaktighet.
- Extern kritik är legitim vad gäller kollektiva skador (miljöpåverkan, risk för desinformation, arbetsförskjutning), men inte vad gäller den interna legitimiteten hos stödet självt.

En särskilt avslöjande dubbelmoral framträder i det utbredda kravet att användning av generativ AI ska avslöjas explicit. Vi kräver inte liknande avslöjande för de flesta andra stöd. Tvärtom firar vi aktivt tekniska framsteg som gör dem osynliga: tjocka glasögon ersatta av kontaktlinser eller refraktiv kirurgi; skrymmande hörapparater miniatyriserade till nästan osynlighet; medicin för fokus, humör eller smärta tagen privat utan fotnot eller friskrivning. I dessa fall behandlar samhället diskret, dold användning som framsteg — som en återställning av värdighet och normalitet. Men när stödet utökar kognition eller uttryck vänder skriptet: nu måste det flaggas, annonseras, motiveras. Osynlighet blir misstänkt istället för önskvärt. Detta selektiva krav på transparens handlar inte verkligen om att förhindra bedrägeri; det handlar om att bevara bekvämligheten med en särskild bild av oasserterat mänskligt författarskap. Fysiska korrigeringar får försvinna; korrigeringar av sinnet måste förbli tydligt märkta.

Om vi ska vara konsekventa måste vi antingen kräva avslöjande för varje stöd (ett absurt och invasivt krav) eller sluta särbehandla kognitiva verktyg för särskild granskning. Den principfasta ståndpunkten — den som respekterar autonomi och värdighet — är att låta varje person besluta hur synligt eller osynligt deras stöd ska vara, utan straffande regler som riktar sig mot en form av hjälp därför att den oroar befintliga uppfattningar om kreativitet och intellekt. Denna essä är inte bara ett försvar för ett särskilt verktyg. Det är ett försvar för den bredare rätten för funktionshindrade och neurodivergenta personer att definiera sina egna tillgångsbehov, utan att behöva motivera dem för de som aldrig har gått i deras skor. Den rätten borde inte vara kontroversiell. Ändå, som den föregående berättelsen visar, är den det fortfarande.